



(12) **EUROPEAN PATENT APPLICATION**

(43) Date of publication:
07.07.2021 Bulletin 2021/27

(21) Application number: **20199347.4**

(22) Date of filing: **30.09.2020**

(51) Int Cl.:
H04N 19/85 ^(2014.01) **G06N 3/04** ^(2006.01)
G03G 15/00 ^(2006.01) **G06F 1/3218** ^(2019.01)
G06F 1/3234 ^(2019.01) **G06F 1/3203** ^(2019.01)
H04N 21/2343 ^(2011.01)

(84) Designated Contracting States:
AL AT BE BG CH CY CZ DE DK EE ES FI FR GB GR HR HU IE IS IT LI LT LU LV MC MK MT NL NO PL PT RO RS SE SI SK SM TR
Designated Extension States:
BA ME KH MA MD TN

(30) Priority: **05.01.2020 US 202062957286 P**
18.01.2020 US 202062962971 P
18.01.2020 US 202062962970 P
09.02.2020 US 202062971994 P
20.04.2020 US 202063012339 P
13.05.2020 US 202063023883 P

(71) Applicant: **Isize Limited**
London E1W 3NX (GB)

(72) Inventors:
• **ANDREOPOULOS, Ioannis**
London, E1W 3NX (GB)
• **GRCE, Srdjan**
London, E1W 3NX (GB)

(74) Representative: **Abel & Imray**
Westpoint Building
James Street West
Bath BA1 2DA (GB)

(54) **PREPROCESSING IMAGE DATA**

(57) A method of preprocessing, prior to encoding with an external encoder, image data using a preprocessing network comprising a set of inter-connected learnable weights is provided. At the preprocessing network, image data from one or more images is received. The image data is processed using the preprocessing network to generate an output pixel representation for en-

coding with the external encoder. The preprocessing network is configured to take as an input display configuration data representing one or more display settings of a display device operable to receive encoded pixel representations from the external encoder. The weights of the preprocessing network are dependent upon the one or more display settings of the display device.

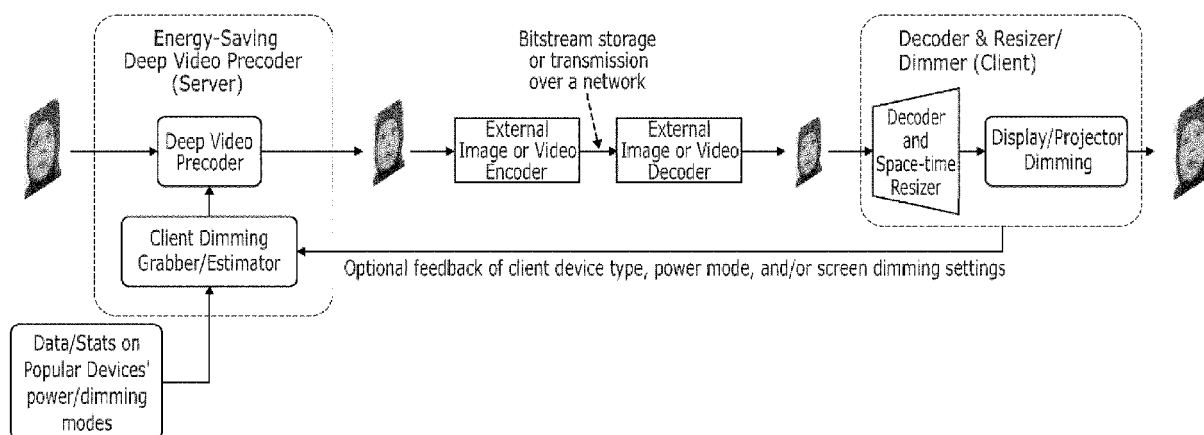


Figure 1

Description

Technical Field

[0001] The present disclosure concerns computer-implemented methods of preprocessing image data prior to encoding with an external encoder. The disclosure is particularly, but not exclusively, applicable where the image data is video data.

Background

[0002] When a set of images or video is sent over a dedicated IP packet switched or circuit-switched connection, a range of streaming and encoding recipes must be selected in order to ensure the best possible use of the available bandwidth. To achieve this, (i) the image or video encoder must be tuned to provide for some bitrate control mechanism; and (ii) the streaming server must provide for the means to control or switch the stream when the bandwidth of the connection does not suffice for the transmitted data. Methods for tackling bitrate control include: constant bitrate (CBR) encoding, variable bitrate (VBR) encoding, or solutions based on a video buffer verifier (VBV) model [9]-[12], such as QVBR, CABR, capped-CRF, etc. These solutions control the parameters of the adaptive quantization and intra-prediction or inter-prediction per image [9]-[12] in order to provide the best possible reconstruction accuracy for the decoded images or video at the smallest number of bits. Methods for tackling stream adaptation are the DASH and HLS protocols - namely, for the case of adaptive streaming over HTTP. Under adaptive streaming, the adaptation comprises the selection of a number of encoding resolutions, bitrates and encoding templates (discussed previously). Therefore, the encoding and streaming process is bound to change the frequency content of the input video and introduce (ideally) imperceptible or (hopefully) controllable quality loss in return for bitrate savings. This quality loss is measured with a range of quality metrics, ranging from low-level signal-to-noise ratio metrics, all the way to complex mixtures of expert metrics that capture higher-level elements of human visual attention and perception. One such metric that is now well-recognised by the video community and the Video Quality Experts Group (VQEG) is the Video Multi-method Assessment Fusion (VMAF), proposed by Netflix. There has been a lot of work in VMAF to make it a "self-interpretable" metric: values close to 100 (e.g. 93 or higher) mean that the compressed content is visually indistinguishable from the original, while low values (e.g. below 70) mean that the compressed content has significant loss of quality in comparison to the original. It has been reported [Ozer, Streaming Media Mag., "Buyers' Guide to Video Quality Metrics", Mar. 29, 2019] that a difference of around 6 points in VMAF corresponds to the so-called Just-Noticeable Difference (JND), i.e. quality difference that will be noticed by the viewer.

[0003] The process of encoding and decoding with a standard image or video encoder always requires the use of linear filters for the production of the decoded (and often upscaled) content that the viewer sees on their device. However, this tends to lead to uncontrolled quality fluctuation in video playback, or poor-quality video playback in general. The viewers most often experience this when they happen to be in an area with poor 4G/WiFi signal strength, where the high-bitrate encoding of a 4K stream will quickly get switched to a much lower-bitrate/lower-resolution encoding, which the decoder and video player will keep on upscaling to the display device's resolution while the viewer continues watching.

[0004] Another factor that affects visual quality significantly is the settings of the decoding device's display or projector (if projection to an external surface or to a mountable or head-mounted display is used). Devices like mobile phones, tablets, monitors, televisions or projectors involve energy-saving modes where the display is deliberately 'dimmed' in order to save power [13][14]. This dimming process tends to include adjustment of contrast and brightness. Such dimming may be adjusted via environmental light via a dedicated sensor, or may be tunable by the viewer or by an external controller, e.g., to reduce eye fatigue [15]. Adaptive streaming proposals have explored energy-driven adaptation, e.g., adjusting transmission parameters or encoding resolutions and frame rates [16]-[21], in order to reduce energy consumption at the client device in a proactive or reactive manner. For example, the work of Massouh et al. [17] proposes luminance preprocessing at the server side in order to reduce power at the client side. The method of Massouh et al. proposes to preprocess the video at the server side in order to reduce its brightness/contrast, thus allowing for energy-consumption reduction at the client side regardless of the energy-saving or dimming settings of the client. However, this reduces the visual quality of the displayed video. In addition, it is attempting to do this experimentally, i.e., via some pre-determined and non-learnable settings.

[0005] Technical solutions to the problem of how to improve visual quality (whilst ensuring efficient operation of encoders) can be grouped into three categories.

[0006] The first type of approaches consists of solutions attempting device-based enhancement, i.e. advancing the state-of-the-art in intelligent video post-processing at the video player. Several of these products are already in the market, including SoC solutions embedded within the latest 8K televisions. While there have been some advances in this domain [21]-[23], this category of solutions is limited by the stringent complexity constraints and power consumption limitations of consumer electronics. In addition, since the received content at the client is already distorted from the compression (quite often severely so), there are theoretical limits to the level of picture detail that can be recovered by client-side post-processing. Further, any additional post-processing taking place in the client side consumes en-

ergy, which goes against the principle of adhering to an energy-saving mode.

[0007] A second family of approaches consists of the development of bespoke image and video encoders, typically based on deep neural networks [16]-[20]. This deviates from encoding, stream-packaging and stream-transport standards and creates bespoke formats, so has the disadvantage of requiring bespoke transport mechanisms and bespoke decoders in the client devices. In addition, in the 50+ years video encoding has been developed most opportunities for improving gain in different situations have been taken, thereby making the current state-of-the-art in spatio-temporal prediction and encoding very difficult to outperform with neural-network solutions that are designed from scratch and learn from data. In addition, all such approaches focus on compression and signal-to-noise ratio (or mean squared error or structural similarity) metrics to quantify distortion, and do not consider the decoding device's dimming settings or how to reverse them to improve quality at the client side and save encoding bitrate.

[0008] The third family of methods comprises perceptual optimisation of existing standards-based encoders by using perceptual metrics during encoding. Here, the challenges are that:

- i) the required tuning is severely constrained by the need for compliance to the utilised standard;
- ii) many of the proposed solutions tend to be limited to focus-of-attention models or shallow learning methods with limited capacity, e.g. assuming that the human gaze is focusing on particular areas of the frame (for instance, in a conversational video we tend to look at the speaker(s), not the background) or using some hand-crafted filters to enhance image slices or groups of image macroblocks prior to encoding;
- iii) such methods tend to require multiple encoding passes, thereby increasing complexity; and
- iv) no consideration is given to the decoding device's display (or projector) settings - i.e., known settings or estimates of those are not exploited in order to improve quality of the displayed video at the decoder side and/or save encoding bitrate for the same quality, as measured objectively by quality metrics or subjectively by human scorers.

[0009] Because of these issues, known designs are very tightly coupled to the specific encoder implementation and ignore the actual decoding display settings. Redesigning them for a new encoder and/or new standard, e.g., from HEVC to VP9 encoding, and/or a wide variety of display configurations, can require substantial effort.

[0010] The present disclosure seeks to solve or mitigate some or all of these above-mentioned problems. Alternatively and/or additionally, aspects of the present disclosure seek to provide improved image and video encoding and decoding methods, and in particular meth-

ods that can be used in combination with existing image and video codec frameworks.

Summary

[0011] In accordance with a first aspect of the present disclosure there is provided a computer-implemented method of preprocessing, prior to encoding with an external encoder, image data using a preprocessing network comprising a set of inter-connected learnable weights, the method comprising:

receiving, at the preprocessing network, image data from one or more images; and
processing the image data using the preprocessing network to generate an output pixel representation for encoding with the external encoder, wherein the preprocessing network is configured to take as an input display configuration data representing one or more display settings of a display device operable to receive encoded pixel representations from the external encoder, and wherein the weights of the preprocessing network are dependent upon the one or more display settings of the display device.

[0012] By conditioning the weights of the preprocessing network on the configuration settings of the external encoder, the representation space can be partitioned within a single model, reducing the need to train multiple models for different device display settings, and reducing the need to redesign and/or reconfigure the preprocessing model for a new display device and/or a display setting. The methods described herein include a preprocessing model that exploits knowledge of display settings (including energy-saving settings) of the display device to tune the parameters and/or operation of the preprocessing model. The preprocessing network may be referred to as an "energy-saving precoder" since it precedes the actual image or video encoder in the image processing pipeline, and aims to optimise the visual quality after decoding and display of image data optionally with dimmed settings, e.g. due to energy-saving features being enabled at the client/display device. This can be done prior to streaming, or during live streaming/transmission of the encoded video, or even in very-low latency video conferencing applications.

[0013] By using a preprocessing network conditioned on the display settings of the display device, the quality of playback at the display device may be improved, by pseudo-reversing the effects of the display device's dimming settings. In particular, an image signal can be enhanced (or amplified) prior to encoding if it is determined that the display device is in an energy-saving mode. This can be achieved without requiring (or imposing) any control on how the display device adjusts its energy-saving or display dimming settings, or indeed how it switches between adaptive streaming modes (eg. in HLS/DASH

or CMAF streaming). Further, a visual quality of the displayed images may be improved for a given encoding bitrate (e.g. as scored by quality metrics) by taking into account client-side display and/or dimming settings. The improvement in visual quality can be achieved whilst causing a minimal increase in energy consumption at the display device, compared to the energy consumption associated with displaying the original video/image without the use of the preprocessing network. Fidelity of the subsequently encoded, decoded and displayed image data to the original image data may also be improved through use of the methods described herein.

[0014] The described methods include technical solutions that are adaptive and/or learnable, which make preprocessing adjustments according to dynamic estimates (or measurements) of client-side display settings, as well as the dynamically-changing characteristics of the content of each image. The methods described herein also use cost functions to optimise the trade-off between signal fidelity and perceptual quality as assessed by human viewers watching the displayed image/video using display devices which support dimming features, e.g. to save energy and/or reduce eye fatigue.

[0015] The methods described herein include a pre-coding method that exploits knowledge, or estimates, of the display settings of the display device to tune the parameters based on data and can utilize a standard image/video encoder with a predetermined encoding recipe for bitrate, quantization and temporal prediction parameters, and fidelity parameters. An overall technical question addressed can be abstracted as: how to optimally preprocess (or "precode") the pixel stream of a video into a (typically) smaller pixel stream, in order to make standards-based encoders as efficient (and fast) as possible and maximise the client-side video quality? This question may be especially relevant where, beyond dimming the display or projector luminance/contrast and other parameters, the client device can resize the content in space and time (e.g. increase/decrease spatial resolution or frame rate) with its existing filters, and/or where perceptual quality is measured with the latest advances in perceptual quality metrics from the literature, e.g., using VMAF or similar metrics.

[0016] In embodiments, the one or more display settings are indicative of an energy-saving state of the display device. For example, the one or more display settings may indicate whether or not the display device is in an energy-saving mode. Different display settings may be used by the display device depending on whether or not the display device is in the energy-saving mode.

[0017] The term 'display device' is used herein to refer to one or multiple physical devices. The display device, or 'client device', can receive encoded image data, decode the encoded image data, and generate an image for display. In some examples, the decoder functionality is provided separately from the display (or projector) functionality, e.g. on different physical devices. In other examples, a single physical device (e.g. a mobile tele-

phone) can both decode encoded bitstreams and display the decoded data on a screen that is part of the device. In some examples, the display device is operable to generate a display output for display on another physical device, e.g. a monitor.

[0018] Preferably, the one or more display settings are indicative of at least one of: whether the display device is plugged into an external power supply or is running on battery power; a battery power level of the display device; voltage or current levels measured or estimated while the display device is decoding and displaying image data; processor utilisation levels of the display device; and a number of concurrent applications or execution threads running on the display device.

[0019] In embodiments, the one or more display settings comprise at least one of: brightness, contrast, gamma correction, refresh rate, flickering settings, bit depth, colour space, colour format, spatial resolution, and back-lighting settings of the display device. Other display settings may be used in alternative embodiments.

[0020] Advantageously, the weights of the preprocessing network are trained using end-to-end back-propagation of errors, the errors calculated using a cost function indicative of an estimated image error associated with displaying output pixel representations using the display device configured according to the one or more display settings.

[0021] In embodiments, the cost function is indicative of an estimate of at least one of: an image noise of the decoded output pixel representation; a bitrate to encode the output pixel representation; and a perceived quality of the decoded output pixel representation. As such, the preprocessing network may be used to reduce noise in the final displayed image(s), reduce the bitrate to encode the output pixel representation, and/or improve the visual quality of the final displayed image(s).

[0022] Advantageously, the cost function is formulated using an adversarial learning framework, in which the preprocessing network is encouraged to generate output pixel representations that reside on the natural image manifold. As such, the preprocessing network is trained to produce images which lie on the natural image manifold, and/or to avoid producing images which do not lie on the natural image manifold. This facilitates an improvement in user perception of the subsequently displayed image(s).

[0023] In embodiments, the estimated image error is indicative of the similarity of the output of decoding the encoded output pixel representation to the received image data based on at least one reference-based quality metric, the at least one reference based quality metric comprising at least one of: an elementwise loss function such as mean squared error, MSE; a structural similarity index metric, SSIM; and a visual information fidelity metric, VIF. As such, the preprocessing network may be used to improve the fidelity of the final displayed image(s) relative to the original input image(s). In embodiments, the weights of the preprocessing network are trained in a

manner that balances perceptual quality of the post-decoded output with fidelity to the original image.

[0024] In embodiments, the resolution of the received image data is different to the resolution of the output pixel representation. For example, the resolution of the output pixel representation may be lower than the resolution of the received image data. By downscaling the image prior to using the external encoder, the external encoder can operate more efficiently by processing a lower resolution image. Moreover, the parameters used when downscaling/upscaling can be chosen to provide different desired results, for example to improve accuracy (i.e. how similarly the recovered images are to the original). Further, the downscaling/upscaling process may be designed to be in accordance with downscaling/upscaling performed by the external encoder, so that the downscaled/upscaled images can be encoded by the external encoder without essential information being lost.

[0025] Preferably, the preprocessing network comprises an artificial neural network including multiple layers having a convolutional architecture, with each layer being configured to receive the output of one or more previous layers.

[0026] In embodiments, the outputs of each layer of the preprocessing network are passed through a non-linear parametric linear rectifier function, pReLU. Other non-linear functions may be used in other embodiments.

[0027] In embodiments, the preprocessing network comprises a dilation operator configured to expand a receptive field of a convolutional operation of a given layer of the preprocessing network. Increasing the receptive field allows for integration of larger global context.

[0028] In embodiments, the weights of the preprocessing network are trained using a regularisation method that controls the capacity of the preprocessing network, the regularisation method comprising using hard or soft constraints and/or a normalisation technique on the weights that reduces a generalisation error.

[0029] In embodiments, the output pixel representation is encoded using the external encoder. In embodiments, the encoded pixel representation is output for transmission, for example to a decoder, for subsequent decoding and display of the image data. In alternative embodiments, the encoded pixel representation is output for storage.

[0030] In accordance with another aspect of the disclosure, there is provided a computer-implemented method of preprocessing one or multiple images into output pixel representations that can subsequently be encoded with any external still-image or video encoder. The method of preprocessing comprises a set of weights inter-connected in a network (termed as "preprocessing network") that ingests: (i) the input pixels from the single or plurality of images; (ii) knowledge or estimates of the client device's display (or projector) settings, or knowledge, or estimates, of the client's power or energy-saving state. These settings can be measurements or estimates of the average settings for a time interval of a video (or

any duration), or can be provided per scene, per individual frame, or even per segment of an individual frame or image.

[0031] In embodiments, the preprocessing network is configured to convert input pixels of each frame to output pixel representations by applying the network weights on the input pixels and accumulating the result of the output product and summation between weights and subsets of input pixels. The network weights, as well as offset or bias terms used for sets of one or more weights, are conditioned on the aforementioned client device power and/or display (or projector) dimming settings. The weights are updated via a training process that uses end-to-end back-propagation of errors computed on the outputs to each group of weights, biases and offsets based on the network connections. The output errors are computed via a cost function that estimates the image or video frame error after encoding and decoding and displaying (or projecting) the output pixel representation of the preprocessing network with the aforementioned external encoder, client decoder and display or projector using encoding and client display settings close to, or identical, to the ones used as inputs to the network.

[0032] In embodiments, the utilized cost function comprises multiple terms that, for the output after decoding, express: image or video frame noise estimates, or functions that estimate the rate to encode the image or video frame, or estimates, functions expressing the perceived quality of the client output from human viewers, or any combinations of these terms. The preprocessing network and cost-function components are trained or refined for any number of iterations prior to deployment (offline) based on training data or, optionally, have their training fine-tuned for any number of iterations based on data obtained during the preprocessing network and encoder-decoder-display operation during deployment (online).

[0033] In embodiments, the disclosed preprocessing network can increase or decrease the resolution of the pixel data in accordance to a given upscaling or downscaling ratio. It can also increase or decrease the frame rate by a given ratio. The ratio can be an integer or fractional number and also includes ratio of 1 (unity) that corresponds to no resolution change. For example, ratio of 2/3 and input image resolution equal to 1080p (1920x1080 pixels, with each pixel comprising 3 color values) would correspond to the output being an image of 720p resolution (1280x720 pixels). Similarly, ratio 2/3 in frame rate and original frame rate of 30 fps would correspond to the output being at 20 fps.

[0034] In terms of the structure of the preprocessing network connections, the network can optionally be structured in a cascaded structure of layers of activations. Each activation in each layer can be connected to any subset (or the entirety) of activations of the next layer, or a subsequent layer by a function determined by the layer weights. In addition, the network can optionally comprise a single or multiple layers of a convolutional architecture, with each layer taking the outputs of the previous layer

and implementing a filtering process via them that realizes the mathematical operation of convolution. In addition, some or all the outputs of each layer can optionally be passed through a non-linear parametric linear rectifier function (pReLU) or other non-linear functions that include, but are not limited to, variations of the sigmoid function or any variation of functions that produce values based on threshold criteria.

[0035] In embodiments, some or all of the convolutional layers of the preprocessing architecture can include implementations of dilation operators that expand the receptive field of the convolutional operation per layer. In addition, the training of the preprocessing network weights can be done with the addition of regularization methods that control the network capacity, via hard or soft constraints or normalization techniques on the layer weights or activations that reduces the generalization error but not the training error.

[0036] In embodiments, the utilized cost functions can express the fidelity to the input images based on reference-based quality metrics that include one or more of:

- (i) elementwise loss functions such as mean squared error (MSE);
- (ii) the sum of absolute errors or variants of it;
- (iii) the structural similarity index metric (SSIM);
- (iv) the visual information fidelity metric (VIF) from the published work of H. Sheikh and A. Bovik entitled "Image Information and Visual Quality",
- (v) the detail loss metric (DLM) from the published work of S. Li, F. Zhang, L. Ma, and K. Ngan entitled "Image Quality Assessment by Separately Evaluating Detail Losses and Additive Impairments";
- (vi) variants and combinations of these metrics;
- (vii) cost functions that express or estimate quality scores attributed to the output images from human viewers;
- (viii) cost functions formulated via an adversarial learning framework, in which the preprocessing network is encouraged to generate output pixel representations that reside on the natural image manifold (and potentially encouraged to reside away from another non-representative manifold); one such example of an adversarial learning framework is the generative adversarial network (GAN), in which the preprocessing network represents the generative component.

[0037] In terms of the provided client side display or projection parameters, these can include: brightness, contrast, gamma correction, refresh rate, flickering (or filtering) settings, bit depth, colour space, colour format, spatial resolution, and back-lighting settings (if existing), whether the device is plugged in an external power supply or is running on battery power, the battery power level, voltage or current levels measured or estimated while the device is decoding and displaying video, CPU or graphics processing unit(s) utilization levels, number of

concurrent applications or execution threads running in the device's task manager or power manager. Moreover, the utilized encoder is a standards-based image or video encoder such as an ISO JPEG or ISO MPEG standard encoder, or a proprietary or royalty-free encoder, such as, but not limited to, an AOMedia encoder.

[0038] In embodiments, either before encoding or after decoding, high resolution and low resolution image or video pairs can be provided and the low resolution image is upsampled and optimized to improve and/or match quality or rate of the high resolution image using the disclosed invention as the means to achieve this. In the optional case of this being applied after decoding, this corresponds to a component on the decoder (client side) that applies such processing after the external decoder has provided the decoded image or video frames.

[0039] In terms of training process across time, the training of the preprocessing network weights, and any adjustment to the cost functions are performed at frequent or infrequent intervals with new measurements from quality, bitrate, perceptual quality scores from humans, or encoded image data from external image or video encoders, and the updated weights and cost functions replace the previously-utilized ones.

[0040] In embodiments, the external encoder comprises an image codec. In embodiments, the image data comprises video data and the one or more images comprise frames of video. In embodiments, the external encoder comprises a video codec.

[0041] The methods of processing image data described herein may be performed on a batch of video data, e.g. a complete video file for a movie or the like, or on a stream of video data.

[0042] The disclosed preprocessing network (or pre-coder) comprises a linear or non-linear system trained by gradient-descent methods and back-propagation. The system can comprise multiple convolutional layers with non-linearities and can include dilated convolutional layers and multiple cost functions, which are used to train the weights of the layers using back propagation. The weights of the layers, their connection and non-linearities, and the associated cost functions used for the training can be conditioned on the image or video encoder settings and the 'dimming' settings of the display device (or estimates of such settings).

[0043] In accordance with another aspect of the disclosure there is provided a computing device comprising:

a processor; and

memory;

wherein the computing device is arranged to perform using the processor any of the methods of preprocessing image data described above.

[0044] In accordance with another aspect of the disclosure there is provided a computer program product arranged, when executed on a computing device comprising a process or memory, to perform any of the meth-

ods of preprocessing image data described above.

[0045] It will of course be appreciated that features described in relation to one aspect of the present disclosure described above may be incorporated into other aspects of the present disclosure.

Description of the Drawings

[0046] Embodiments of the present disclosure will now be described by way of example only with reference to the accompanying schematic drawings of which:

Figure 1 is a schematic diagram of a method of processing image data in accordance with embodiments;

Figures 2(a) to 2(d) are schematic diagrams showing a preprocessing network in accordance with embodiments;

Figure 3 is a schematic diagram showing a preprocessing network in accordance with embodiments;

Figure 4 is a schematic diagram showing a convolutional layer of a preprocessing network in accordance with embodiments;

Figure 5 is a schematic diagram showing a training process in accordance with embodiments;

Figure 6 is an example of video playback in accordance with embodiments;

Figure 7 is a flowchart showing the steps of a method of preprocessing image data in accordance with embodiments; and

Figure 8 is a schematic diagram of a computing device in accordance with embodiments.

Detailed Description

[0047] Embodiments of the present disclosure are now described.

[0048] Figure 1 is a schematic diagram showing a method of processing image data, according to embodiments. The method involves deep video precoding conditional to knowledge or estimates of a client device's dimming settings (and optional resizing). Image or video input data is pre-processed by a conditional 'precoder' (conditioned to information or estimates provided by the client dimming grabber/estimator) prior to passing to an external image or video codec. Information on the client dimming or power modes can be obtained via feedback from the client, as shown in Figure 1. This can be achieved, for example, via a Javascript component running on a client browser page, which will send real-time information on whether the device is on power-saving mode or it is plugged into a charger. Other mechanisms include, but are not limited to, a player application on the client side that has permissions to read the detailed status of the display and power modes from the device. If such feedback is not available from the client, as shown in Figure 1, offline estimates and statistics can be used for the power modes and display dimming settings of popular

devices, e.g., mobile phones, tablets, smart TVs, etc. The user can even select to switch to a dimming-optimized precoding mode manually, by selecting the appropriate option during the video playback. Embodiments are applicable to batch processing, i.e. processing a group of images or video frames together without delay constraints (e.g. an entire video sequence), as well as to stream processing, i.e. processing only a limited subset of a stream of images or video frames, or even a select subset of a single image, due to delay or buffering constraints.

[0049] As discussed above, embodiments comprise a deep conditional precoding that processes input image or video frames. The deep conditional precoding (and optional post-processing) shown in Figure 1 can comprise any combination of learnable weights locally or globally connected in a network with a non-linear activation function. An example of such weights is shown in Figure 2(a) and an associated example in Figure 2(b) showcases global connectivity between weights and inputs. That is, Figure 2(a) shows a combination of inputs $x_0, \dots, x_3$ with weight coefficients θ and linear or non-linear activation function $g()$, and Figure 2(b) is a schematic diagram showing layers of interconnected activations and weights, forming an artificial neural network with global connectivity. An example of local connectivity between weights and inputs is shown in Figure 2(c) for a 2D dilated convolution [1], 3×3 kernel, and dilation rate of 2. As such, Figure 2(c) is a schematic diagram of a 2D dilated convolutional layer with local connectivity. Figure 2(d) is a schematic diagram of back-propagation of errors δ from an intermediate layer (right hand side of Figure 2(d)) to the previous intermediate layer using gradient descent.

[0050] An example of the deep conditional precoding model is shown in Figure 3. It consists of a series of conditional convolutional layers and elementwise parametric ReLu (pReLu) layers of weights and activations. As such, Figure 3 shows a cascade of conditional convolutional and parametric ReLu (pReLu) layers mapping input pixel groups to transformed output pixel groups. In embodiments, all layers receive codec settings as input, along with the representation from the previous layer. In alternative embodiments, some layers do not receive codec settings as input. There is also an optional skip connection between the input and output layer. In embodiments, each conditional convolution takes the output of the preceding layer as input (with the first layer receiving the image as input), along with intended (or estimated) settings for the client side display or projection, encoded as a numerical representation. For image precoding for an image codec, these settings can include but are not limited to: (i) the provided or estimated decoding device's display settings include at least one of the following: brightness, contrast, gamma correction, refresh rate, flickering (or filtering) settings, bit depth, colour space, colour format, spatial resolution, and back-lighting settings (if existing); (ii) the provided or estimated decoding device's power or energy-saving settings include at least

one of the following: whether the device is plugged in an external power supply or is running on battery power, the battery power level, voltage or current levels measured or estimated while the device is decoding and displaying video, CPU or graphics processing unit(s) utilization levels, number of concurrent applications or execution threads running in the device's task manager or power manager.

[0051] Figure 4 is a schematic diagram showing a convolutional layer conditional to client device dimmer and/or energy settings. The layer receives these settings (or their estimates) as an input, which is quantized and one-hot encoded. The one-hot encoded vector is then mapped to intermediate representations w and b , which respectively weight and bias the channels of the output of the dilated convolutional layer z . The one-hot encoded conditional convolutional layer is shown in this example for the case of JPEG encoding. In the example for conditional convolutional layers for JPEG encoding of Figure 4, these settings are quantized to integer values and one-hot encoded. The one-hot encoding is then mapped via linear or non-linear functions, such as densely connected layers (e.g. as shown in Figure 2(b)), to vector representations. These vector representations are then used to weight and bias the output of a dilated convolution - thus conditioning the dilated convolution on the user settings.

[0052] Conditioning the precoding on user settings enables a partitioning of the representation space within a single model without having to train multiple models for every possible client-side display or power setting. An example of the connectivity per dilated convolution is as shown in Figure 2(c). The dilation rate (spacing between each learnable weight in the kernel), kernel size (number of learnable weights in the kernel per dimension) and stride (step per dimension in the convolution operation) are all variable per layer, with a dilation rate of 1 equating to a standard convolution. Increasing the dilation rate increases the receptive field per layer and allows for integration of larger global context. The entirety of the series of dilated convolutional layers and activation functions can be trained end-to-end based on back-propagation of errors for the output layer backwards using gradient descent methods, as illustrated in Figure 2(d).

[0053] An example of the framework for training the deep conditional precoding is shown in Figure 5. In particular, Figure 5 is a schematic diagram showing training of deep conditional precoding for intra-frame coding, where s represents the scale factor for resizing and Q represents the client device display settings (which may include energy settings), or their estimates. The discriminator and precoder are trained iteratively and the perceptual model can also be trained iteratively with the precoder, or pre-trained and frozen. The guidance image input $\hat{x}$ to the discriminator refers to a linearly down-scaled, compressed and up-scaled representation of x . The post-processing refers to a simple linear (non-parametric) upscaling in this example. In embodiments, the precoding is trained in a manner that balances perceptual

quality of the post-decoded and displayed output with fidelity to the original image or frame. The precoding is trained iteratively via backpropagation and any variation of gradient descent, e.g. as described with reference to Figure 2(d). Parameters of the learning process, such as the learning rate, the use of dropout and other regularization options to stabilize the training and convergence process are applied.

[0054] The example training framework described herein assumes that post-processing only constitutes a simple linear resizing. The framework comprises a linear or non-linear weighted combination of loss functions for training the deep conditional precoding. The loss functions will now be described.

[0055] The distortion loss $\mathcal{L}_D$ is derived as a function of a perceptual model, and optimized over the precoder weights, in order to match or maximize the perceptual quality of the post-decoded output $\hat{x}$ over the original input x . The perceptual model is a parametric model that estimates the perceptual quality of the post-decoded output $\hat{x}$. The perceptual model can be configured as an artificial neural network with weights and activation functions and connectivity similar to those shown in Figures 2(a)-2(d). This perceptual model produces a reference or non-reference based score for quality; reference based scores compare the quality of $\hat{x}$ to x , whereas non-reference based scores produce a blind image quality assessment of $\hat{x}$. The perceptual model can optionally approximate non-differentiable perceptual score functions, including VIF, ADM2 and VMAF, with continuous differentiable functions. The perceptual model can also be trained to output human rater scores, including MOS or distributions over ACR values. Figure 5 depicts a non-reference based example framework trained to output the distribution over ACR values. The perceptual model can either be pre-trained or trained iteratively with the deep conditional

precoding by minimizing perceptual loss $\mathcal{L}_P$ and $\mathcal{L}_D$ alternately or sequentially respectively. The perceptual

loss $\mathcal{L}_P$ is a function of the difference between the reference (human-rater) quality scores and model-predicted quality scores over a range of inputs. The distortion

loss $\mathcal{L}_D$ can thus be defined between $\hat{x}$ and x , as a linear or non-linear function of the intermediate activations of selected layers of the perceptual model, up to the output reference or non-reference based scores. Additionally, in order to ensure faithful reconstruction of the input x , the distortion loss is combined with a pixel-wise loss directly between the input x and $\hat{x}$, such as mean absolute error (MAE) or mean squared error (MSE), and optionally a structural similarity loss, based on SSIM or MSSIM.

[0056] The adversarial loss $\mathcal{L}_A$ is optimized over the precoder weights, in order to ensure that the post-decoded output $\hat{x}$, which is generated via the precoder, lies on the natural image manifold. The adversarial loss is for-

ulated by modelling the precoder as a generator and adding a discriminator into the framework, which corresponds to the generative adversarial network (GAN) setup [2], as illustrated in Figure 5. In the standard GAN configuration, the discriminator receives the original input frames, represented by x and the post-decoded output $\hat{x}$ as input, which can respectively be referred to as 'real' and 'fake' (or 'artificial') data. The discriminator is trained to distinguish between the 'real' and 'fake' data with loss $\mathcal{L}_C$. On the contrary, the precoder is trained with $\mathcal{L}_A$ to fool the discriminator into classifying the 'fake' data as 'real'. The discriminator and precoder are trained alternately with $\mathcal{L}_C$ and $\mathcal{L}_A$ respectively, with additional constraints such as gradient clipping depending on the GAN variant. The loss formulations for $\mathcal{L}_C$ and $\mathcal{L}_A$ directly depend on the GAN variant utilized; this can include but is not limited to standard saturating, non-saturating [2],[3] and least-squares GANs [4] and their relativistic GAN counterparts [5], and integral probability metric (IPM) based GANs, such as Wasserstein GAN (WGAN) [6],[7] and Fisher GAN [8]. Additionally, the loss functions can be patch-based (i.e. evaluated between local patches of x and $\hat{x}$) or can be image-based (i.e. evaluated between whole images). The discriminator is configured with conditional convolutional layers, e.g. as shown in Figure 3 and Figure 4. An additional guidance image or frame $\tilde{x}$ is passed to the discriminator, which can represent a linear downsampled, upsampled and compressed representation of x , following the same scaling and codec settings as $\hat{x}$. The discriminator can thus learn to distinguish between x , $\hat{x}$ and $\tilde{x}$, whilst the precoder can learn to generate representations that post-decoding and scaling will be perceptually closer to x than $\tilde{x}$.

[0057] The noise loss component $\mathcal{L}_N$ is optimized over the precoder weights and acts as a form of regularization, in order to further ensure that the precoder is trained such that the post-decoded and displayed output is a denoised representation of the input. Examples of noise include aliasing artefacts (e.g. jaggging or ringing) introduced by downscaling in the precoder, as well as additional codec artefacts (e.g. blocking and brightness/contrast/gamma adjustment) introduced by the virtual codec and display during training to emulate a standard video or image codec that performs lossy compression and a display or projector that performs dimming functionalities. An example of the noise loss component $\mathcal{L}_N$ is total variation denoising, which is effective at removing noise while preserving edges.

[0058] The rate loss $\mathcal{L}_R$ is an optional loss component that is optimized over the precoder weights, in order to constrain the rate (number of bits or bitrate) of the precoder output, as estimated by a virtual codec module.

[0059] The virtual codec and virtual display modules depicted in Figure 5 emulate: (i) a standard image or

video codec that performs lossy compression and primarily consists of a frequency transform component, a quantization and entropy encoding component and a dequantization and inverse transform component; (ii) a standard display or projector that adjusts brightness/contrast/gamma/refresh rate/flicker parameters, etc. The virtual codec module takes as input the precoder output and any associated codec settings (e.g. CRF, preset). The virtual display module takes as input the client dimmer/energy saving settings or estimates that the precoder itself is conditioned on (e.g. via the conditional convolutional layers shown in Figure 4). The frequency transform component of the virtual codec is any variant of discrete sine or cosine transform or wavelet transform, or an atom-based decomposition. The dequantization and inverse transform component can convert the transform coefficients back into approximated pixel values. The main source of loss for the virtual codec module comes from the quantization component, which emulates any multi-stage deadzone or non-deadzone quantizer. The virtual display module applies standard contrast/brightness/gamma/flicker filters and optionally frame and resolution up/down conversion using differentiable approximations of these functionalities. Any non-differentiable parts of the standard codec or standard display functionalities are approximated with continuous differentiable alternatives; one such example is the rounding operation in quantization, which can be approximated with additive uniform noise of support width equal to 1. In this way, the entire virtual codec and virtual display modules are end-to-end

continuously differentiable. To estimate the rate in $\mathcal{L}_R$, the entropy coding component represents a continuously differentiable approximation to a standard Huffman, arithmetic or runlength encoder, or any combination of those that is also made context adaptive, i.e., by looking at quantization symbol types and surrounding values (context conditioning) in order to utilize the appropriate probability model and compression method. The entropy coding and other virtual codec components can be made learnable, with an artificial neural network or similar, and jointly trained with the precoding or pre-trained to maximize the likelihood on the frequency transformed and quantized precoder representations. Alternatively, a given lossy JPEG, MPEG or AOMedia open encoder can be used to provide the actual rate and compressed representations as reference, which the virtual codec can be trained to replicate. In both cases, training of the artificial neural network parameters can be performed with backpropagation and gradient descent methods.

[0060] As shown in the example depicted in Figure 5, the discriminator and deep conditional precoding can be trained alternately. This can also be true for the perceptual model and deep video precoding (or otherwise the perceptual model can be pre-trained and weights frozen throughout precoder training). After training one component, its weights are updated and it is frozen and the other component is trained. This weight update and in-

terleaved training improves both and allows for end-to-end training and iterative improvement both during the training phase. The number of iterations that a component is trained before being frozen, $n \geq 1$. For the discriminator-precoding pair, this will depend on the GAN loss formulation and whether one seeks to train the discriminator to optimality. Furthermore, the training can continue online and at any time during the system's operation. An example of this is when new images and quality scores are added into the system, or new forms of display dimming options and settings are added, which correspond to a new or updated form of dimming options, or new types of image content, e.g., cartoon images, images from computer games, virtual or augmented reality applications, etc.

[0061] To test the methods described herein, a utilized video codec fully-compliant to the H.264/AVC standard was used, with the source code being the JM19.0 reference software of the HHI/Fraunhofer repository [29]. For the experiments, the same encoding parameters were used, which were: encoding frame rate of 25 frames-per-second; YUV encoding with zero U, V channels since the given images are monochrome (zero-valued UV channels consume minimal bitrate that is equal for both the described methods and the original video encoder); one I frame (only first); motion estimation search range ± 32 pixels and simplified UMHExagon search selected; 2 reference frames; and P prediction modes enabled (and B prediction modes enabled for QP-based control); NumberBFrames parameter set to 0 for rate-control version and NumberBFrames set to 3 for QP control version; CABAC is enabled and single-pass encoding is used; single-slice encoding (no rate sacrificed for error resilience); in the rate-control version, InitialQP=32 and all default rate control parameters of the encoder.cfg file of JM19.0 were enabled; SourceBitDepthLuma/ Chroma set to 12 bits and no use of rescaling or Q-Matrix.

[0062] The source material comprised standard 1080p 8-bit RGB resolution videos, but similar results have been obtained with visual image sequences or videos in full HD or ultra-HD resolution and any dynamic range for the input pixel representations. For the display dimming functionalities, standard brightness and contrast downconversions were applied by 60% when the utilized mobile, tablet, monitor or smart TV is set on energy-saving mode. These settings were communicated to the disclosed network preprocessing system as shown in Figure 1. A conditional neural network architecture based on the examples shown in Figures 2(a)-2(d) and Figure 3 was used to implement the conditional precoding system. The training and testing followed the examples described with reference to Figure 4 and Figure 5. An indicative visual result is shown in Figure 6, which depicts an example of video playback of a sequence where the left side of the video has not been adjusted using the precoder, and the right side has been adjusted using the methods disclosed herein, to compensate for screen dimming and improve the visual quality. As shown by Figure 6, for the provided

video sequence and under the knowledge or approximation of the encoding and display dimming parameters, the described methods offer significant quality improvement by pseudo-reversing the effects of encoding and dimming. This occurs for both types of encodings (bitrate and QP control). Beyond the presented embodiments, the methods described herein can be realized with the full range of options and adaptivity described in the previous examples, and all such options and their adaptations are covered by this disclosure.

[0063] Using as an option selective downscaling during the precoding process and allowing for a linear up-scaling component at the client side after decoding (as presented in Figure 1), the methods described herein can shrink the input to 10%-40% of the frame size of the input frames, which means that the encoder processes a substantially smaller number of pixels and is therefore 2-6 times faster than the encoder of the full resolution infrared image sequence. This offers additional benefits in terms of increased energy autonomy for video monitoring under battery support, vehicle/mobile/airborne visual monitoring systems, etc.

[0064] Figure 7 shows a method 700 for preprocessing image data using a preprocessing network comprising a set of inter-connected learnable weights. The method 700 may be performed by a computing device, according to embodiments. The method 700 may be performed at least in part by hardware and/or software. The preprocessing is performed prior to encoding the preprocessed image data with an external encoder. The preprocessing network is configured to take as an input display configuration data representing one or more display settings of a display device operable to receive encoded pixel representations from the external encoder. The weights of the preprocessing network are dependent upon (i.e. conditioned on) the one or more display settings of the display device. At item 710, image data from one or more images is received at the preprocessing network. The image data may be retrieved from storage (e.g. in a memory), or may be received from another entity. At item 720, the image data is processed using the preprocessing network (e.g. by applying the weights of the preprocessing network to the image data) to generate an output pixel representation for encoding with the external encoder. In embodiments, the method 700 comprises encoding the output pixel representation, e.g. using the external encoder. The encoded output pixel representation may be transmitted, for example to the display device for decoding and subsequent display.

[0065] Embodiments of the disclosure include the methods described above performed on a computing device, such as the computing device 800 shown in Figure 8. The computing device 800 comprises a data interface 801, through which data can be sent or received, for example over a network. The computing device 800 further comprises a processor 802 in communication with the data interface 801, and memory 803 in communication with the processor 802. In this way, the computing device

800 can receive data, such as image data or video data, via the data interface 801, and the processor 802 can store the received data in the memory 803, and process it so as to perform the methods of described herein, including preprocessing image data prior to encoding using an external encoder, and optionally encoding the pre-processed image data.

[0066] Each device, module, component, machine or function as described in relation to any of the examples described herein may comprise a processor and/or processing system or may be comprised in apparatus comprising a processor and/or processing system. One or more aspects of the embodiments described herein comprise processes performed by apparatus. In some examples, the apparatus comprises one or more processing systems or processors configured to carry out these processes. In this regard, embodiments may be implemented at least in part by computer software stored in (non-transitory) memory and executable by the processor, or by hardware, or by a combination of tangibly stored software and hardware (and tangibly stored firmware). Embodiments also extend to computer programs, particularly computer programs on or in a carrier, adapted for putting the above described embodiments into practice. The program may be in the form of non-transitory source code, object code, or in any other non-transitory form suitable for use in the implementation of processes according to embodiments. The carrier may be any entity or device capable of carrying the program, such as a RAM, a ROM, or an optical memory device, etc.

[0067] Various measures (including methods, apparatus, computing devices and computer program products) are provided for preprocessing of a single or a plurality of images prior to encoding them with an external image or video encoder. The preprocessing method comprises a set of weights, biases and offset terms that are interconnected (termed as "preprocessing network") that ingests: (i) the input pixels from the single or plurality of images; (ii) knowledge, or estimates, of the decoder device's display (or projector) settings, or knowledge, or estimates, of the decoder device's power or energy-saving state. The utilised preprocessing network is configured to convert input pixels to an output pixel representation such that: weights and offset or bias terms of the network are conditioned on the aforementioned display or projection or decoder power or energy-saving settings and the weights are trained end-to-end with back-propagation of errors from outputs to inputs. The output errors are computed via a cost function that estimates the image or video frame error after encoding, decoding and displaying the output pixel representation of the preprocessing network with the aforementioned external encoder and decoder, as well as a display or projector that uses settings similar to, or identical, to the ones used as inputs to the network. The utilized cost function comprises multiple terms that, for the output after decoding, express: image or video frame noise estimates, or functions or training data that estimate the rate to encode the image

or video frame, or estimates, functions or training data expressing the perceived quality of the output from human viewers, or any combinations of these terms.

[0068] In embodiments, the resolution or frame rate of the pixel data is increased or decreased in accordance to a given increase or decrease ratio and also includes ratio of 1 (unity) that corresponds to no resolution or frame rate change.

[0069] In embodiments, weights in the preprocessing network are used, in order to construct a function of the input over single or multiple layers of a convolutional architecture, with each layer receiving outputs of the previous layers.

[0070] In embodiments, the outputs of each layer of the preprocessing network are passed through a non-linear parametric linear rectifier function (pReLU) or other non-linear activation function.

[0071] In embodiments, the convolutional layers of the preprocessing architecture include dilation operators that expand the receptive field of the convolutional operation per layer.

[0072] In embodiments, the training of the preprocessing network weights is done with the addition of regularization methods that control the network capacity, via hard or soft constraints or normalization techniques on the layer weights or activations that reduces the generalization error.

[0073] In embodiments, cost functions are used that express the fidelity to the input images based on reference-based quality metrics that include one or more of: elementwise loss functions such as mean squared error (MSE); a structural similarity index metric (SSIM); a visual information fidelity metric (VIF), for example from the published work of H. Sheikh and A. Bovik entitled "Image Information and Visual Quality"; a detail loss metric (DLM), for example from the published work of S. Li, F. Zhang, L. Ma, and K. Ngan entitled "Image Quality Assessment by Separately Evaluating Detail Losses and Additive Impairments"; variants and combinations of these metrics.

[0074] In embodiments, cost functions are used that express or estimate quality scores attributed to the output images from human viewers.

[0075] In embodiments, cost functions are used that are formulated via an adversarial learning framework, in which the preprocessing network is encouraged to generate output pixel representations that reside on the natural image manifold (and optionally encouraged to reside away from another non-representative manifold).

[0076] In embodiments, the provided or estimated decoding device's display settings include at least one of the following: brightness, contrast, gamma correction, refresh rate, flickering (or filtering) settings, bit depth, colour space, colour format, spatial resolution, and back-lighting settings (if existing). In embodiments, the provided or estimated decoding device's power or energy-saving settings include at least one of the following: whether the device is plugged in an external power supply or is run-

ning on battery power, the battery power level, voltage or current levels measured or estimated while the device is decoding and displaying video, CPU or graphics processing unit(s) utilization levels, number of concurrent applications or execution threads running in the device's task manager or power manager.

[0077] In embodiments, the utilized encoder is a standards-based image or video encoder such as an ISO JPEG or ISO MPEG standard encoder, or a proprietary or royalty-free encoder, such as, but not limited to, an AOMedia encoder.

[0078] In embodiments, high resolution and low resolution image or video pairs are provided and the low resolution image is upsampled and optimized to improve and/or match quality or rate to the high resolution image.

[0079] In embodiments, the training of the preprocessing network weights and any adjustment to the cost functions are performed at frequent or infrequent intervals with new measurements from decoder display settings, perceptual quality scores from humans, or encoded image data from external image or video encoders, and the updated weights and cost functions replace the previously-utilized ones.

[0080] While the present disclosure has been described and illustrated with reference to particular embodiments, it will be appreciated by those of ordinary skill in the art that the disclosure lends itself to many different variations not specifically illustrated herein.

[0081] Where in the foregoing description, integers or elements are mentioned which have known, obvious or foreseeable equivalents, then such equivalents are herein incorporated as if individually set forth. Reference should be made to the claims for determining the true scope of the present invention, which should be construed so as to encompass any such equivalents. It will also be appreciated by the reader that integers or features of the disclosure that are described as preferable, advantageous, convenient or the like are optional and do not limit the scope of the independent claims. Moreover, it is to be understood that such optional integers or features, whilst of possible benefit in some embodiments of the disclosure, may not be desirable, and may therefore be absent, in other embodiments.

References

[0082]

- [1] F. Yu and V. Koltun, "Multi-scale context aggregation by dilated convolutions," arXiv preprint arXiv: 1511.07122, 2015.
- [2] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville and Y. Bengio, "Generative adversarial nets," in Advances in neural information processing systems, 2014.
- [3] T. Salimans, I. Goodfellow, W. Zaremba, V. Cheung, A. Radford and X. Chen, "Improved techniques for training gans," in Advances in neural information

processing systems, 2016.

[4] X. Mao, Q. Li, H. Xie, R. Y. K. Lau, Z. Wang and S. Paul Smolley, "Least squares generative adversarial networks," in Proceedings of the IEEE International Conference on Computer Vision, 2017.

[5] A. Jolicoeur-Martineau, "The relativistic discriminator: a key element missing from standard GAN," arXiv preprint arXiv: 1807.00734, 2018.

[6] M. Arjovsky, S. Chintala and L. Bottou, "Wasserstein gan," arXiv preprint arXiv: 1701.07875, 2017.

[7] I. Gulrajani, F. Ahmed, M. Arjovsky, V. Dumoulin and A. C. Courville, "Improved training of wasserstein gans," in Advances in neural information processing systems, 2017.

[8] Y. Mroueh and T. Sercu, "Fisher gan," in Advances in Neural Information Processing Systems, 2017.

[9] Boyce, Jill, et al. "Techniques for layered video encoding and decoding." U.S. Patent Application No. 13/738,138.

[10] Dar, Yehuda, and Alfred M. Bruckstein. "Improving low bit-rate video coding using spatio-temporal down-scaling." arXiv preprint arXiv: 1404.4026 (2014).

[11] Martemyanov, Alexey, et al. "Real-time video coding/decoding." U.S. Patent No. 7,336,720. 26 Feb. 2008.

[12] van der Schaar, Mihaela, and Mahesh Balakrishnan. "Spatial scalability for fine granular video encoding." U.S. Patent No. 6,836,512. 28 Dec. 2004.

[13] Hayashi et al., "Dimmer and video display device using the same", US Patent 10,078,236 B2, Date of patent: Sept. 18, 2018.

[14] Ato et al., "Display device," US Patent 9,791,701 B2, Date of patent: Oct. 17, 2017.

[15] Jung, "Liquid crystal display with brightness extractor and driving method thereof for modulating image brightness by controlling the average picture level to reduce glare and eye fatigue," US Patent 8,970,635 B2, Date of patent: Mar. 3, 2015.

[16] Varghese, Benoy, et al. "e-DASH: Modelling an energy-aware DASH player." 2017 IEEE 18th International Symposium on A World of Wireless, Proc. IEEE Mobile and Multimedia Networks (WoWMoM), 2017.

[17] Massouh, Nizar, et al. "Experimental study on luminance preprocessing for energy-aware HTTP-based mobile video streaming." Proc. IEEE 2014 5th European Workshop on Visual Information Processing (EUVIP).

[18] Hu, Wenjie, and Guohong Cao. "Energy-aware video streaming on smartphones." Proc. IEEE Conf. on Computer Communications (INFOCOM). IEEE, 2015.

[19] Almowuena, Saleh, et al. "Energy-aware and bandwidth-efficient hybrid video streaming over mobile networks." IEEE Trans. on Multimedia 18.1 (2015): 102-115.

- [20] Mehrabi, Abbas, et al. "Energy-aware QoE and backhaul traffic optimization in green edge adaptive mobile video streaming." IEEE Trans. on Green Communications and Networking 3.3 (2019): 828-839.
- [21] Dong, Jie, and Yan Ye. "Adaptive downsampling for high-definition video coding." IEEE Transactions on Circuits and Systems for Video Technology 24.3 (2014): 480-488.
- [22] Douma, Peter, and Motoyuki Koike. "Method and apparatus for video upscaling." U.S. Patent No. 8,165,197. 24 Apr. 2012.
- [23] Su, Guan-Ming, et al. "Guided image up-sampling in video coding." U.S. Patent No. 9,100,660. 4 Aug. 2015.
- [24] Hinton, Geoffrey E., and Ruslan R. Salakhutdinov. "Reducing the dimensionality of data with neural networks." science313.5786 (2006): 504-507.
- [25] van den Oord, Aaron, et al. "Conditional image generation with pixelcnn decoders." Advances in Neural Information Processing Systems. 2016.
- [26] Theis, Lucas, et al. "Lossy image compression with compressive autoencoders." arXiv preprint arXiv: 1703.00395(2017).
- [27] Wu, Chao-Yuan, Nayan Singhal, and Philipp Krähenbühl. "Video Compression through Image Interpolation." arXiv preprint arXiv: 1804.06919 (2018).
- [28] Rippel, Oren, and Lubomir Bourdev. "Real-time adaptive image compression." arXiv preprint arXiv: 1705.05823 (2017).
- [29] K. Suehring, HHI AVC reference code repository, online at the HHI website.
- [30] G. Bjontegaard, "Calculation of average PSNR differences between RD-curves," VCEG-M33 (2001).

Claims

1. A computer-implemented method of preprocessing, prior to encoding with an external encoder, image data using a preprocessing network comprising a set of inter-connected learnable weights, the method comprising:

receiving, at the preprocessing network, image data from one or more images; and
 processing the image data using the preprocessing network to generate an output pixel representation for encoding with the external encoder,
 wherein the preprocessing network is configured to take as an input display configuration data representing one or more display settings of a display device operable to receive encoded pixel representations from the external encoder, and

wherein the weights of the preprocessing network are dependent upon the one or more display settings of the display device.

2. A method according to claim 1, wherein the one or more display settings are indicative of an energy-saving state of the display device.
3. A method according to claim 2, wherein the one or more display settings are indicative of at least one of: whether the display device is plugged into an external power supply or is running on battery power; a battery power level of the display device; voltage or current levels measured or estimated while the display device is decoding and displaying image data; processor utilisation levels of the display device; and a number of concurrent applications or execution threads running on the display device.
4. A method according to any preceding claim, wherein the one or more display settings comprise at least one of: brightness, contrast, gamma correction, refresh rate, flickering settings, bit depth, colour space, colour format, spatial resolution, and back-lighting settings of the display device.
5. A method according to any preceding claim, comprising training the weights of the preprocessing network using end-to-end back-propagation of errors, the errors calculated using a cost function indicative of an estimated image error associated with displaying output pixel representations using the display device configured according to the one or more display settings.
6. A method according to claim 5, wherein the cost function is indicative of an estimate of at least one of:

an image noise of the decoded output pixel representation;
 a bitrate to encode the output pixel representation; and
 a perceived quality of the decoded output pixel representation.

7. A method according to claim 5 or claim 6, wherein the cost function is formulated using an adversarial learning framework, in which the preprocessing network is encouraged to generate output pixel representations that reside on the natural image manifold.
8. A method according to any of claims 5 to 7, wherein the estimated image error is indicative of the similarity of the output of decoding the encoded output pixel representation to the received image data based on at least one reference-based quality metric, the at least one reference based quality metric comprising at least one of:

an elementwise loss function such as mean squared error, MSE;
 a structural similarity index metric, SSIM; and
 a visual information fidelity metric, VIF.

5

9. A method according to any preceding claim, wherein the resolution of the received image data is different to the resolution of the output pixel representation.

10. A method according to any preceding claim, wherein the preprocessing network comprises an artificial neural network including multiple layers having a convolutional architecture, with each layer being configured to receive the output of one or more previous layers.

10

15

11. A method according to claim 10, comprising passing the outputs of each layer of the preprocessing network through a non-linear parametric linear rectifier function, pReLU.

20

12. A method according to any preceding claim, wherein the preprocessing network comprises a dilation operator configured to expand a receptive field of a convolutional operation of a given layer of the preprocessing network.

25

13. A method according to any preceding claim, wherein the weights of the preprocessing network are trained using a regularisation method that controls the capacity of the preprocessing network, the regularisation method comprising using hard or soft constraints and/or a normalisation technique on the weights that reduces a generalisation error.

30

35

14. A computing device comprising:

a processor; and
 memory;

wherein the computing device is arranged to perform, using the processor, a method of preprocessing image data according to any of claims 1 to 13.

40

15. A computer program product arranged, when executed on a computing device comprising a processor or memory, to perform a method of preprocessing image data according to any of claims 1 to 13.

45

50

55

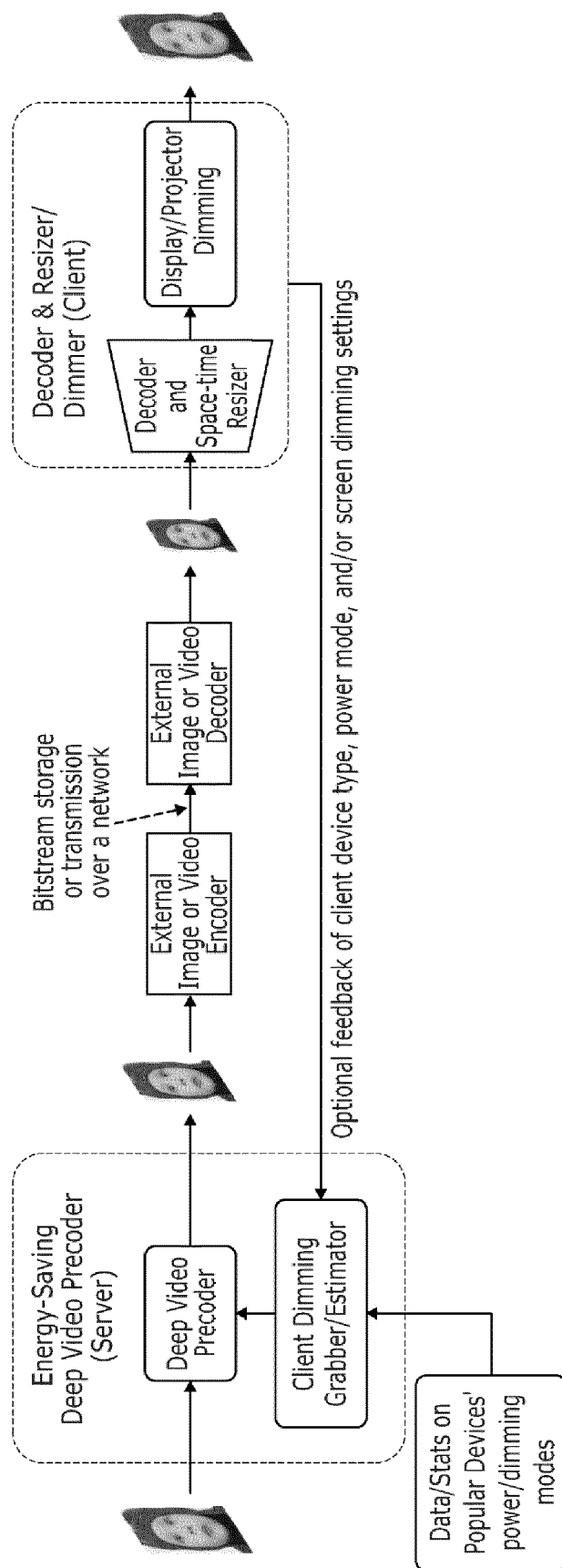


Figure 1

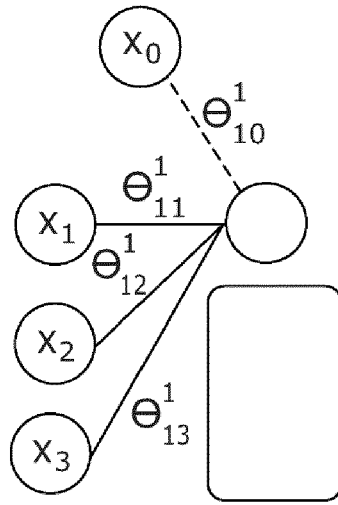


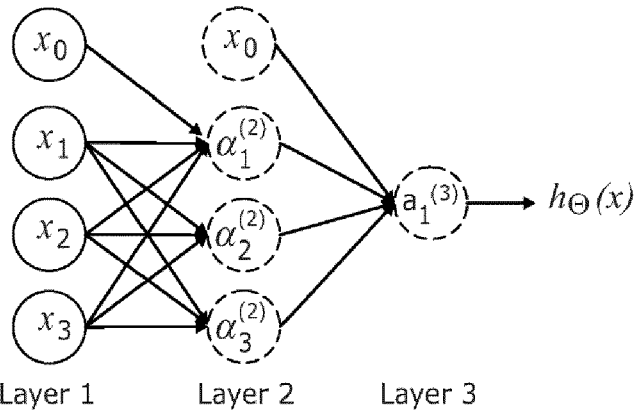
Figure 2(a)

$$g(\Theta_{10}^1 x_0 + \Theta_{11}^1 x_1 + \Theta_{12}^1 x_2 + \Theta_{13}^1 x_3)$$

Θ_{13}^1 means:

- 1 (superscript) - mapping from layer 1
- 1 - mapping to node 1 in layer 2 (L+1)
- 3 - mapping from node 3 in layer 1 (L)

Figure 2(b)



$$\alpha_1^{(2)} = g(\Theta_{10}^{(1)} x_0 + \Theta_{11}^{(1)} x_1 + \Theta_{12}^{(1)} x_2 + \Theta_{13}^{(1)} x_3)$$

$$\alpha_2^{(2)} = g(\Theta_{20}^{(1)} x_0 + \Theta_{21}^{(1)} x_1 + \Theta_{22}^{(1)} x_2 + \Theta_{23}^{(1)} x_3)$$

$$\alpha_3^{(2)} = g(\Theta_{30}^{(1)} x_0 + \Theta_{31}^{(1)} x_1 + \Theta_{32}^{(1)} x_2 + \Theta_{33}^{(1)} x_3)$$

$$h_{\Theta}(x) = \alpha_1^{(3)} = g(\Theta_{10}^{(2)} \alpha_0^{(2)} + \Theta_{11}^{(2)} \alpha_1^{(2)} + \Theta_{12}^{(2)} \alpha_2^{(2)} + \Theta_{13}^{(2)} \alpha_3^{(2)})$$

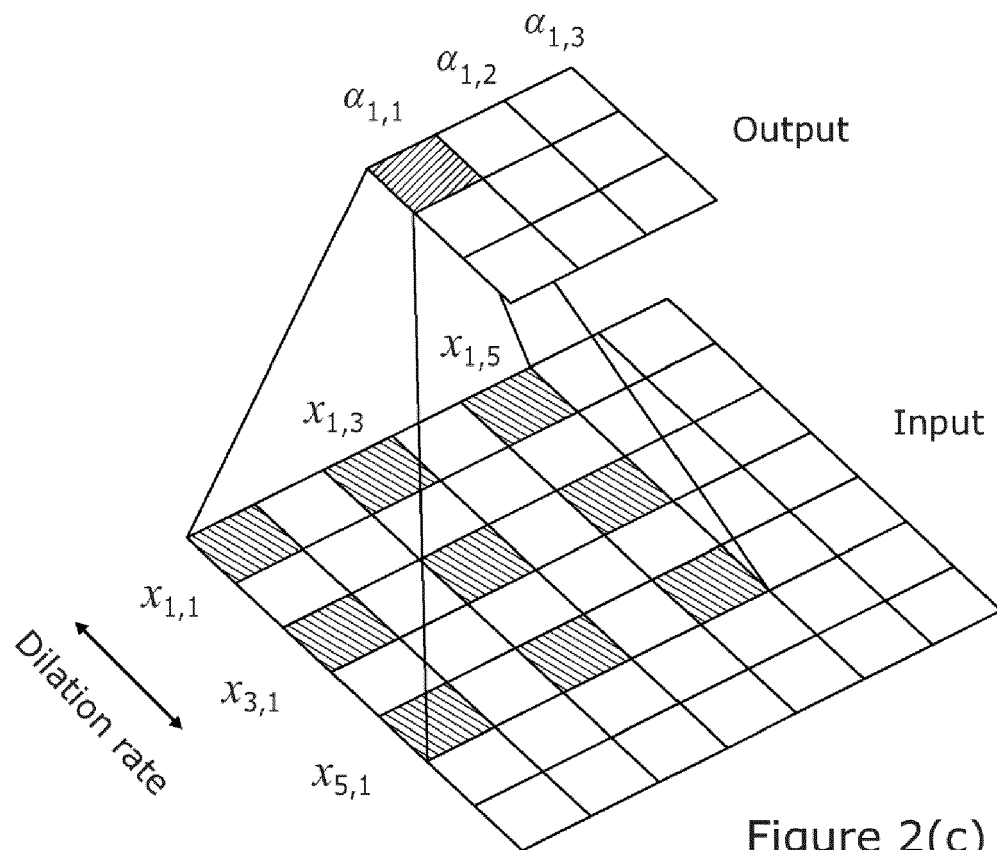


Figure 2(c)

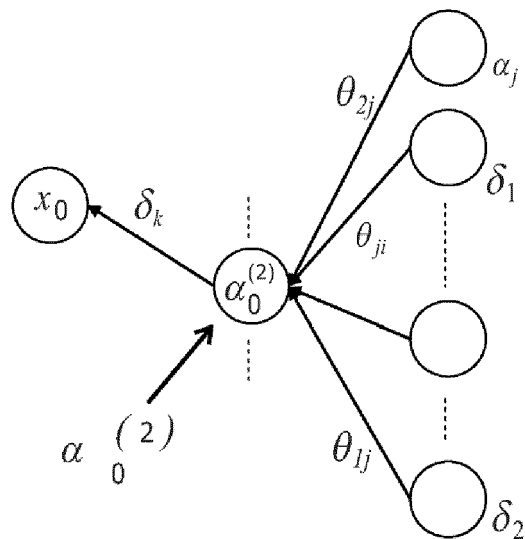


Figure 2(d)

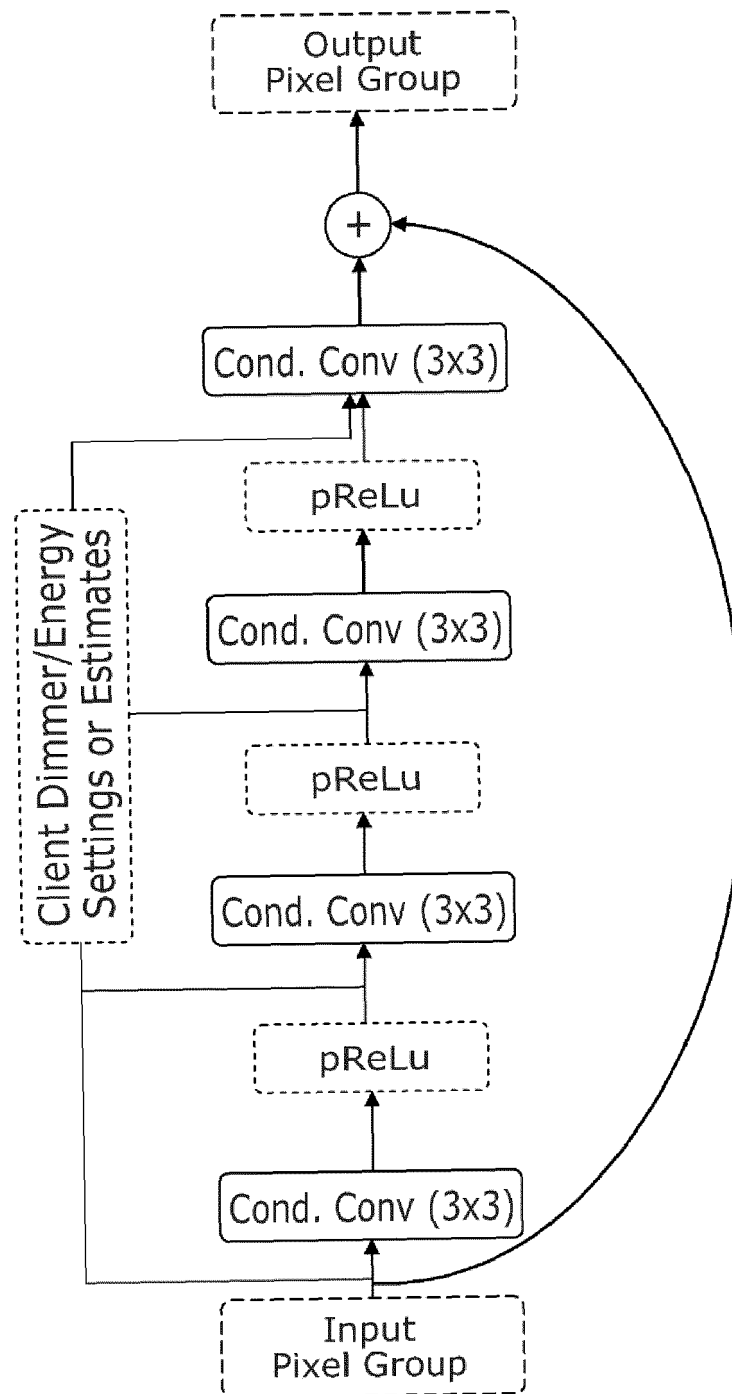


Figure 3

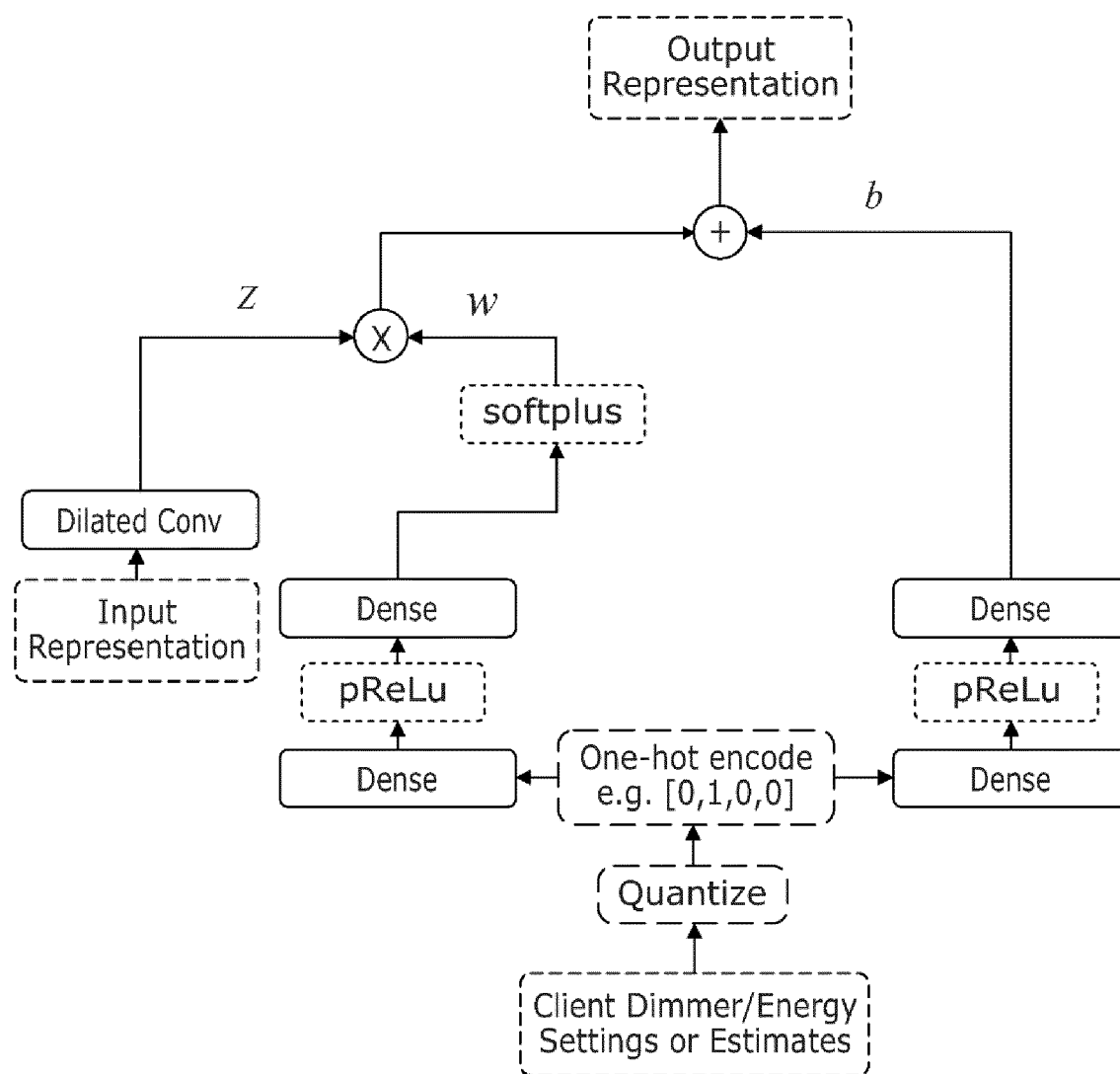


Figure 4

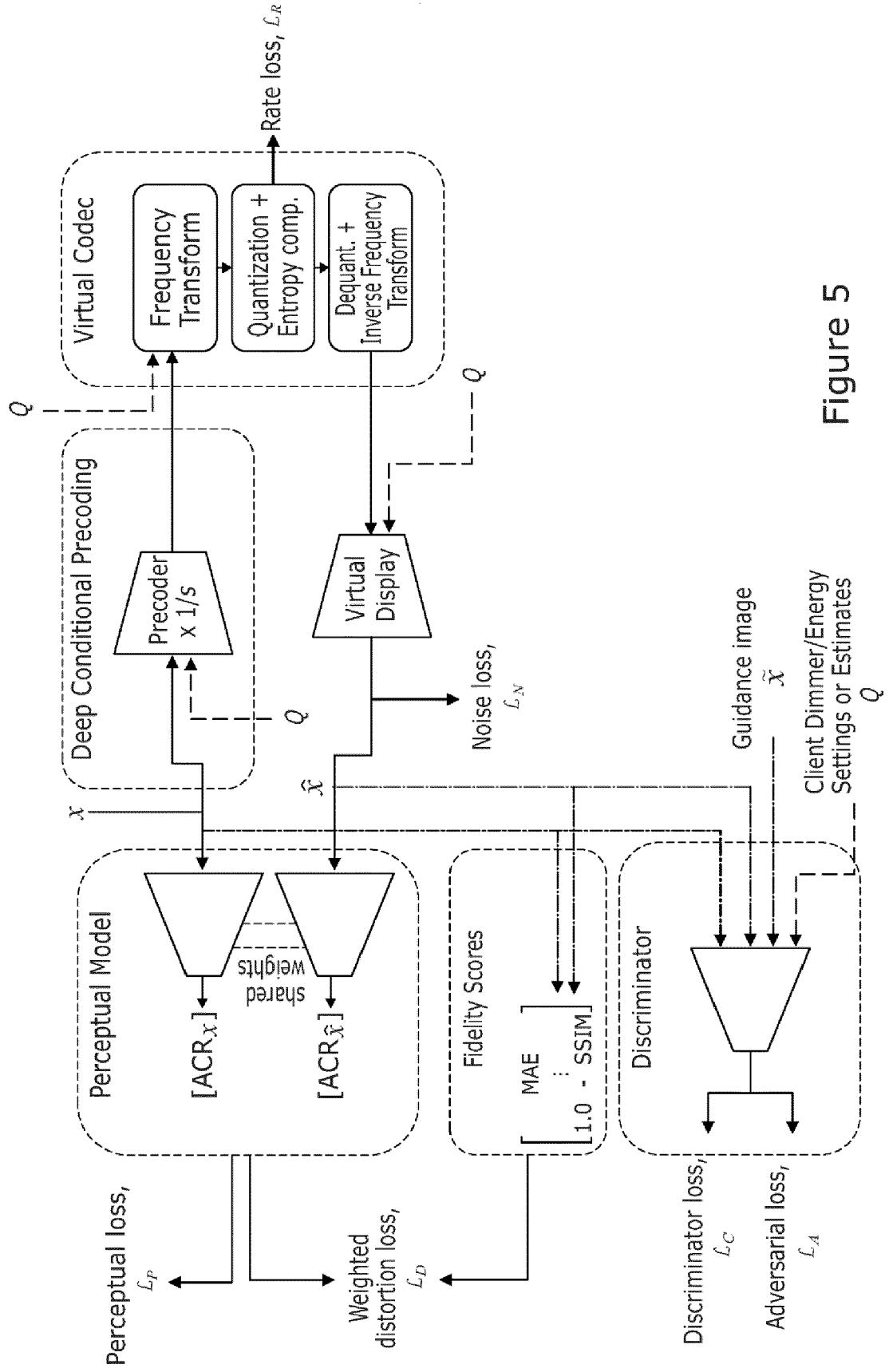


Figure 5

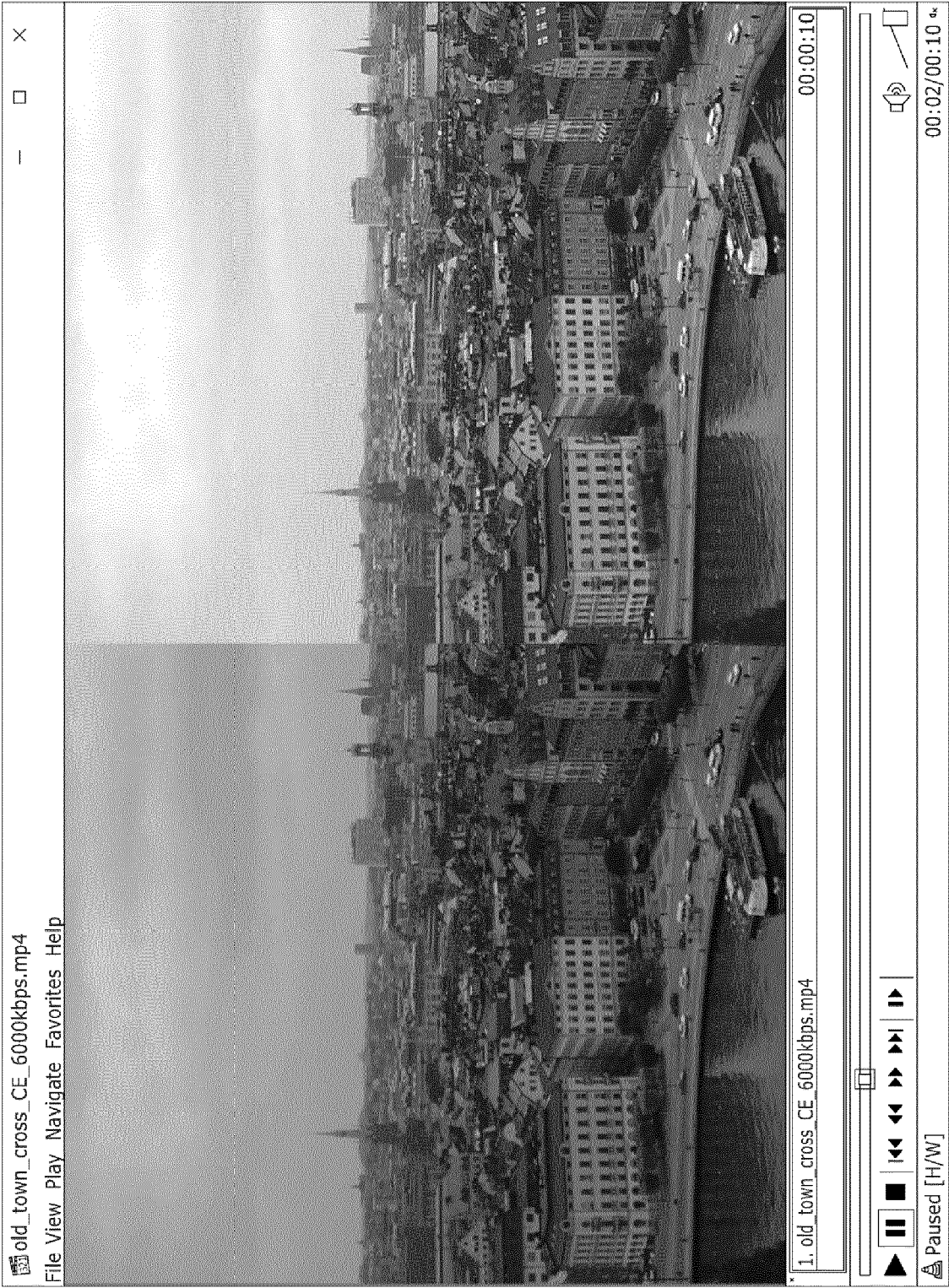


Figure 6

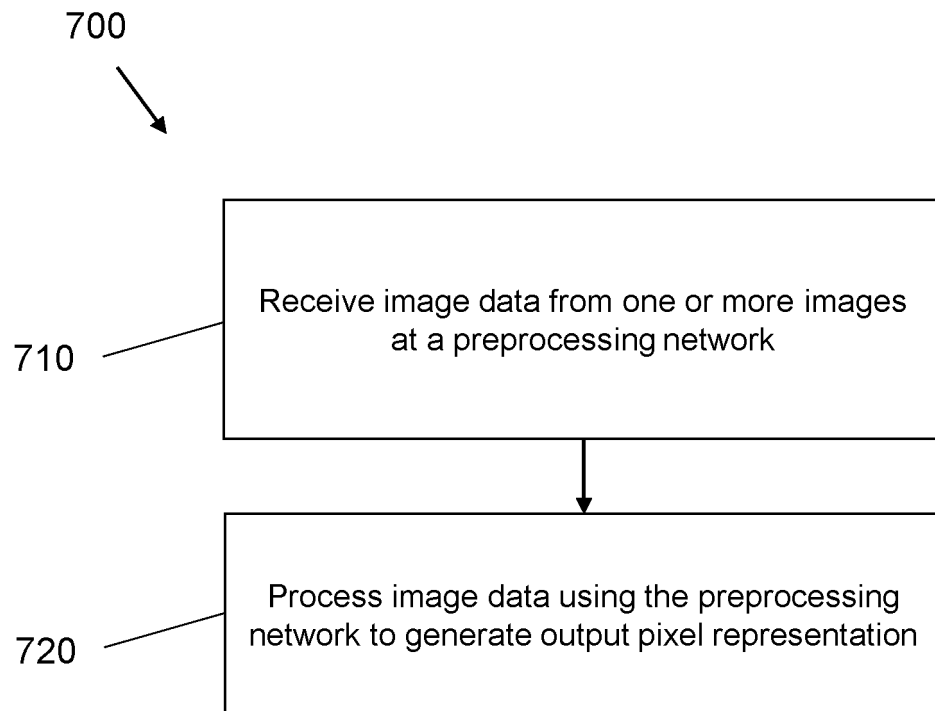


Figure 7

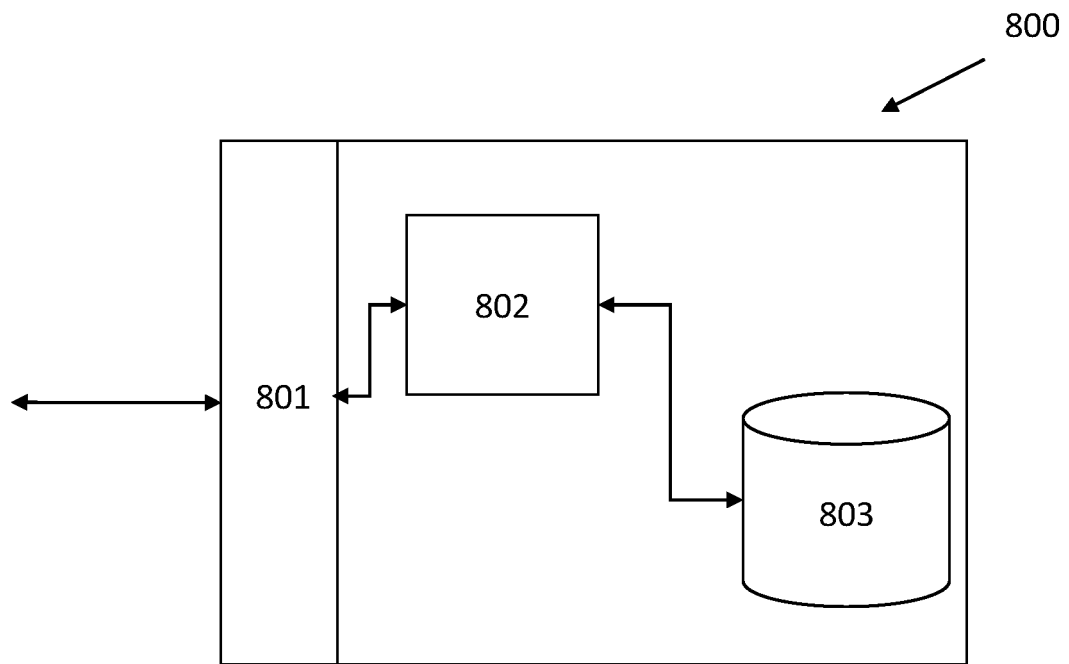


Figure 8



EUROPEAN SEARCH REPORT

Application Number
EP 20 19 9347

5

10

15

20

25

30

35

40

45

50

55

DOCUMENTS CONSIDERED TO BE RELEVANT			
Category	Citation of document with indication, where appropriate, of relevant passages	Relevant to claim	CLASSIFICATION OF THE APPLICATION (IPC)
X	WO 2019/009449 A1 (SAMSUNG ELECTRONICS CO LTD [KR]) 10 January 2019 (2019-01-10)	1,5,6,8,10,11,14,15	INV. H04N19/85 G06N3/04 G03G15/00 G06F1/3218 G06F1/3234 G06F1/3203 H04N21/2343
Y	* paragraphs [0066], [0067], [0073], [0082], [0086] - [0092], [0098], [0115], [0135], [0196] * * figures 2, 4B, 5B, 6, 7A * & US 2020/145661 A1 (JEON SUN-YOUNG [KR] ET AL) 7 May 2020 (2020-05-07)	2-4,7,9,12,13	
Y	EP 3 584 676 A1 (GUANGDONG OPPO MOBILE TELECOMMUNICATIONS CORP LTD [CN]) 25 December 2019 (2019-12-25) * paragraphs [0018], [0084] - [0091] * * figure 7 *	2-4	
Y	GB 2 548 749 A (MAGIC PONY TECH LTD [GB]) 27 September 2017 (2017-09-27) * page 8, line 5 - line 9 * * page 83, line 2 - line 6 * * page 69, line 32 - page 70, line 2 *	7,9	
A		1-6,8,10-15	TECHNICAL FIELDS SEARCHED (IPC)
Y	CN 110 248 190 A (UNIV XI AN JIAOTONG) 17 September 2019 (2019-09-17) * paragraph [0028] *	12	H04N G06N G03G G06F
Y	WO 2018/229490 A1 (UCL BUSINESS PLC [GB]) 20 December 2018 (2018-12-20) * paragraphs [0223] - [0226] * & US 2020/167930 A1 (WANG GUOTAI [GB] ET AL) 28 May 2020 (2020-05-28)	13	
A	WO 2019/197712 A1 (NOKIA TECHNOLOGIES OY [FI]) 17 October 2019 (2019-10-17) * paragraph [0150] *	1-15	
	-/--		
The present search report has been drawn up for all claims			
Place of search The Hague		Date of completion of the search 23 October 2020	Examiner Votini, Stefano
CATEGORY OF CITED DOCUMENTS X : particularly relevant if taken alone Y : particularly relevant if combined with another document of the same category A : technological background O : non-written disclosure P : intermediate document		T : theory or principle underlying the invention E : earlier patent document, but published on, or after the filing date D : document cited in the application L : document cited for other reasons & : member of the same patent family, corresponding document	

EPO FORM 1503 03.82 (P04C01)



EUROPEAN SEARCH REPORT

Application Number
EP 20 19 9347

5

10

15

20

25

30

35

40

45

50

55

DOCUMENTS CONSIDERED TO BE RELEVANT			
Category	Citation of document with indication, where appropriate, of relevant passages	Relevant to claim	CLASSIFICATION OF THE APPLICATION (IPC)
A	US 2009/304071 A1 (SHI XIAOJIN [US] ET AL) 10 December 2009 (2009-12-10) * paragraph [0050] *	1-15	
A	US 2012/033040 A1 (PAHALAWATTA PESHALA V [US] ET AL) 9 February 2012 (2012-02-09) * paragraphs [0028], [0029], [0033], [0034], [0035], [0039], [0040]; figures 3, 4, 6, 7 *	1-15	
A	Song Han ET AL: "Learning both Weights and Connections for Efficient Neural Networks", 30 October 2015 (2015-10-30), XP055396330, Retrieved from the Internet: URL:https://arxiv.org/pdf/1506.02626.pdf [retrieved on 2017-08-04] * paragraph [0003] * * figures 2-3 *	1-15	
A	WO 2016/132152 A1 (MAGIC PONY TECH LTD [GB]; WANG ZEHAN [GB] ET AL.) 25 August 2016 (2016-08-25) * figures 1-27 *	1-15	
A	DENG YUBIN ET AL: "Image Aesthetic Assessment: An experimental survey", IEEE SIGNAL PROCES SING MAGAZINE, IEEE SERVICE CENTER, PISCATAWAY, NJ, US, vol. 34, no. 4, 1 July 2017 (2017-07-01), pages 80-106, XP011656215, ISSN: 1053-5888, DOI: 10.1109/MSP.2017.2696576 [retrieved on 2017-07-11] * abstract *	1-15	
The present search report has been drawn up for all claims			TECHNICAL FIELDS SEARCHED (IPC)
Place of search The Hague		Date of completion of the search 23 October 2020	Examiner Votini, Stefano
CATEGORY OF CITED DOCUMENTS X : particularly relevant if taken alone Y : particularly relevant if combined with another document of the same category A : technological background O : non-written disclosure P : intermediate document		T : theory or principle underlying the invention E : earlier patent document, but published on, or after the filing date D : document cited in the application L : document cited for other reasons & : member of the same patent family, corresponding document	

EPO FORM 1503 03.82 (P04C01)

ANNEX TO THE EUROPEAN SEARCH REPORT
ON EUROPEAN PATENT APPLICATION NO.

EP 20 19 9347

5

This annex lists the patent family members relating to the patent documents cited in the above-mentioned European search report. The members are as contained in the European Patent Office EDP file on
The European Patent Office is in no way liable for these particulars which are merely given for the purpose of information.

23-10-2020

10

15

20

25

30

35

40

45

50

55

Patent document cited in search report	Publication date	Patent family member(s)	Publication date
WO 2019009449 A1	10-01-2019	CN 110832860 A	21-02-2020
		EP 3624452 A1	18-03-2020
		KR 20200016879 A	17-02-2020
		US 2020145661 A1	07-05-2020
		WO 2019009449 A1	10-01-2019
		WO 2019009489 A1	10-01-2019
EP 3584676 A1	25-12-2019	CN 106933326 A	07-07-2017
		EP 3584676 A1	25-12-2019
		US 2020008153 A1	02-01-2020
		WO 2018161571 A1	13-09-2018
GB 2548749 A	27-09-2017	NONE	
CN 110248190 A	17-09-2019	NONE	
WO 2018229490 A1	20-12-2018	EP 3639240 A1	22-04-2020
		US 2020167930 A1	28-05-2020
		WO 2018229490 A1	20-12-2018
WO 2019197712 A1	17-10-2019	NONE	
US 2009304071 A1	10-12-2009	NONE	
US 2012033040 A1	09-02-2012	CN 102450009 A	09-05-2012
		CN 104954789 A	30-09-2015
		DK 2663076 T3	05-12-2016
		EP 2422521 A1	29-02-2012
		EP 2663076 A2	13-11-2013
		ES 2602326 T3	20-02-2017
		HK 1214440 A1	22-07-2016
		JP 5044057 B2	10-10-2012
		JP 5364820 B2	11-12-2013
		JP 2012231526 A	22-11-2012
		JP 2012521734 A	13-09-2012
		US 2012033040 A1	09-02-2012
		WO 2010123855 A1	28-10-2010
WO 2016132152 A1	25-08-2016	EP 3259910 A1	27-12-2017
		EP 3259911 A1	27-12-2017
		EP 3259912 A1	27-12-2017
		EP 3259913 A1	27-12-2017
		EP 3259914 A1	27-12-2017
		EP 3259915 A1	27-12-2017
		EP 3259916 A1	27-12-2017
		EP 3259918 A1	27-12-2017

EPO FORM P0459

For more details about this annex : see Official Journal of the European Patent Office, No. 12/82

**ANNEX TO THE EUROPEAN SEARCH REPORT
ON EUROPEAN PATENT APPLICATION NO.**

EP 20 19 9347

5

This annex lists the patent family members relating to the patent documents cited in the above-mentioned European search report.
The members are as contained in the European Patent Office EDP file on
The European Patent Office is in no way liable for these particulars which are merely given for the purpose of information.

23-10-2020

10

15

20

25

30

35

40

45

50

55

Patent document cited in search report	Publication date	Patent family member(s)	Publication date
		EP 3259919 A1	27-12-2017
		EP 3259920 A1	27-12-2017
		EP 3493149 A1	05-06-2019
		GB 2539845 A	28-12-2016
		GB 2539846 A	28-12-2016
		GB 2540889 A	01-02-2017
		GB 2543429 A	19-04-2017
		GB 2543958 A	03-05-2017
		US 2017345130 A1	30-11-2017
		US 2017347060 A1	30-11-2017
		US 2017347061 A1	30-11-2017
		US 2017347110 A1	30-11-2017
		US 2017374374 A1	28-12-2017
		US 2018129918 A1	10-05-2018
		US 2018130177 A1	10-05-2018
		US 2018130178 A1	10-05-2018
		US 2018130179 A1	10-05-2018
		US 2018130180 A1	10-05-2018
		WO 2016132145 A1	25-08-2016
		WO 2016132146 A1	25-08-2016
		WO 2016132147 A1	25-08-2016
		WO 2016132148 A1	25-08-2016
		WO 2016132149 A1	25-08-2016
		WO 2016132150 A1	25-08-2016
		WO 2016132151 A1	25-08-2016
		WO 2016132152 A1	25-08-2016
		WO 2016132153 A1	25-08-2016
		WO 2016132154 A1	25-08-2016

EPO FORM P0459

For more details about this annex : see Official Journal of the European Patent Office, No. 12/82

REFERENCES CITED IN THE DESCRIPTION

This list of references cited by the applicant is for the reader's convenience only. It does not form part of the European patent document. Even though great care has been taken in compiling the references, errors or omissions cannot be excluded and the EPO disclaims all liability in this regard.

Patent documents cited in the description

- US 738138 [0082]
- US 7336720 B [0082]
- US 6836512 B [0082]
- US 10078236 B2 [0082]
- US 9791701 B2 [0082]
- US 8970635 B2 [0082]
- US 8165197 B [0082]
- US 9100660 B [0082]

Non-patent literature cited in the description

- **OZER.** Buyers' Guide to Video Quality Metrics. *Streaming Media Mag.*, 29 March 2019 [0002]
- **H. SHEIKH ; A. BOVIK.** *Image Information and Visual Quality* [0073]
- **S. LI ; F. ZHANG ; L. MA ; K. NGAN.** *Image Quality Assessment by Separately Evaluating Detail Losses and Additive Impairments* [0073]
- **F. YU ; V. KOLTUN.** Multi-scale context aggregation by dilated convolutions. *arXiv preprint arXiv: 1511.07122*, 2015 [0082]
- **I. GOODFELLOW ; J. POUGET-ABADIE ; M. MIRZA ; B. XU ; D. WARDE-FARLEY ; S. OZAIR ; A. COURVILLE ; Y. BENGIO.** Generative adversarial nets. *Advances in neural information processing systems*, 2014 [0082]
- **T. SALIMANS ; I. GOODFELLOW ; W. ZAREMBA ; V. CHEUNG ; A. RADFORD ; X. CHEN.** Improved techniques for training gans. *Advances in neural information processing systems*, 2016 [0082]
- **X. MAO ; Q. LI ; H. XIE ; R. Y. K. LAU ; Z. WANG ; S. PAUL SMOLLEY.** Least squares generative adversarial networks. *Proceedings of the IEEE International Conference on Computer Vision*, 2017 [0082]
- **A. JOLICOEUR-MARTINEAU.** The relativistic discriminator: a key element missing from standard GAN. *arXiv preprint arXiv: 1807.00734*, 2018 [0082]
- **M. ARJOVSKY ; S. CHINTALA ; L. BOTTOU.** Wasserstein gan. *arXiv preprint arXiv: 1701.07875*, 2017 [0082]
- **I. GULRAJANI ; F. AHMED ; M. ARJOVSKY ; V. DUMOULIN ; A. C. COURVILLE.** Improved training of wasserstein gans. *Advances in neural information processing systems*, 2017 [0082]
- **Y. MROUEH ; T. SERCU.** Fisher gan. *Advances in Neural Information Processing Systems*, 2017 [0082]
- **DAR, YEHUDA ; ALFRED M. BRUCKSTEIN.** Improving low bit-rate video coding using spatio-temporal down-scaling. *arXiv preprint arXiv: 1404.4026*, 2014 [0082]
- **VARGHESE, BENOY et al.** e-DASH: Modelling an energy-aware DASH player. *2017 IEEE 18th International Symposium on A World of Wireless, Proc. IEEE Mobile and Multimedia Networks (WoWMoM)*, 2017 [0082]
- **MASSOUH, NIZAR et al.** Experimental study on luminance preprocessing for energy-aware HT-TP-based mobile video streaming. *Proc. IEEE 2014 5th European Workshop on Visual Information Processing (EUVIP)* [0082]
- Energy-aware video streaming on smartphones. **HU, WENJIE ; GUOHONG CAO.** *Proc. IEEE Conf. on Computer Communications (INFOCOM)*. IEEE, 2015 [0082]
- **ALMOWUENA, SALEH et al.** Energy-aware and bandwidth-efficient hybrid video streaming over mobile networks. *IEEE Trans. on Multimedia*, 2015, vol. 18.1, 102-115 [0082]
- **MEHRABI, ABBAS et al.** Energy-aware QoE and backhaul traffic optimization in green edge adaptive mobile video streaming. *IEEE Trans. on Green Communications and Networking*, 2019, vol. 3.3, 828-839 [0082]
- **DONG, JIE ; YAN YE.** Adaptive downsampling for high-definition video coding. *IEEE Transactions on Circuits and Systems for Video Technology*, 2014, vol. 24.3, 480-488 [0082]
- **HINTON, GEOFFREY E. ; RUSLAN R. SALAKHUTDINOV.** Reducing the dimensionality of data with neural networks. *science*, 2006, vol. 313.5786 (504-507) [0082]
- **VAN DEN OORD, AARON et al.** Conditional image generation with pixelcnn decoders. *Advances in Neural Information Processing Systems*, 2016 [0082]
- **THEIS, LUCAS et al.** Lossy image compression with compressive autoencoders. *arXiv preprint arXiv: 1703.00395*, 2017 [0082]
- **WU, CHAO-YUAN ; NAYAN SINGHAL ; PHILIPP KRÄHENBÜHL.** Video Compression through Image Interpolation. *arXiv preprint arXiv: 1804.06919*, 2018 [0082]

EP 3 846 477 A1

- **RIPPEL, OREN ; LUBOMIR BOURDEV.** Real-time adaptive image compression. *arXiv preprint arXiv:1705.05823*, 2017 [0082]
- **K. SUEHRING.** HHI AVC reference code repository. *HHI website* [0082]
- **G. BJONTEGAARD.** Calculation of average PSNR differences between RD-curves. *VCEG-M33*, 2001 [0082]